

FD²: A Dedicated Framework for Fine-Grained Dataset Distillation

Hongxu Ma^{1†}, Guang Li^{2†*}, Shijie Wang³, Dongzhan Zhou⁴, Baoli Sun⁵,
Takahiro Ogawa², Miki Haseyama², and Zhihui Wang^{5*}

¹ Zhejiang University

² Hokkaido University

³ The University of Queensland

⁴ Shanghai AI Laboratory

⁵ Dalian University of Technology

[†] Equal Contribution, ^{*} Corresponding Authors:

guang@lmd.ist.hokudai.ac.jp, zhihuiwang@dlut.edu.cn

Abstract. Dataset distillation (DD) compresses a large training set into a small synthetic set, reducing storage and training cost, and has shown strong results on general benchmarks. Decoupled DD further improves efficiency by splitting the pipeline into pretraining, sample distillation, and soft-label generation. However, existing decoupled methods largely rely on coarse class-label supervision and optimize samples within each class in a nearly identical manner. On fine-grained datasets, this often yields distilled samples that (i) retain large intra-class variation with subtle inter-class differences and (ii) become overly similar within the same class, limiting localized discriminative cues and hurting recognition. To solve the above-mentioned problems, we propose **FD²**, a dedicated framework for **F**ine-grained **D**ataset **D**istillation. FD² localizes discriminative regions and constructs fine-grained representations for distillation. During pretraining, counterfactual attention learning aggregates discriminative representations to update class prototypes. During distillation, a fine-grained characteristic constraint aligns each sample with its class prototype while repelling others, and a similarity constraint diversifies attention across same-class samples. Experiments on multiple fine-grained and general datasets show that FD² integrates seamlessly with decoupled DD and improves performance in most settings, indicating strong transferability. Code is available at <https://github.com/Guang000/FD2>.

Keywords: Dataset Distillation · Fine-grained · Decoupled Distillation

1 Introduction

Recently, the rapid growth of large-scale datasets has driven substantial progress in artificial intelligence, but it also brings high computational cost and long training time [7, 8, 12]. Dataset distillation (DD) addresses this challenge by learning

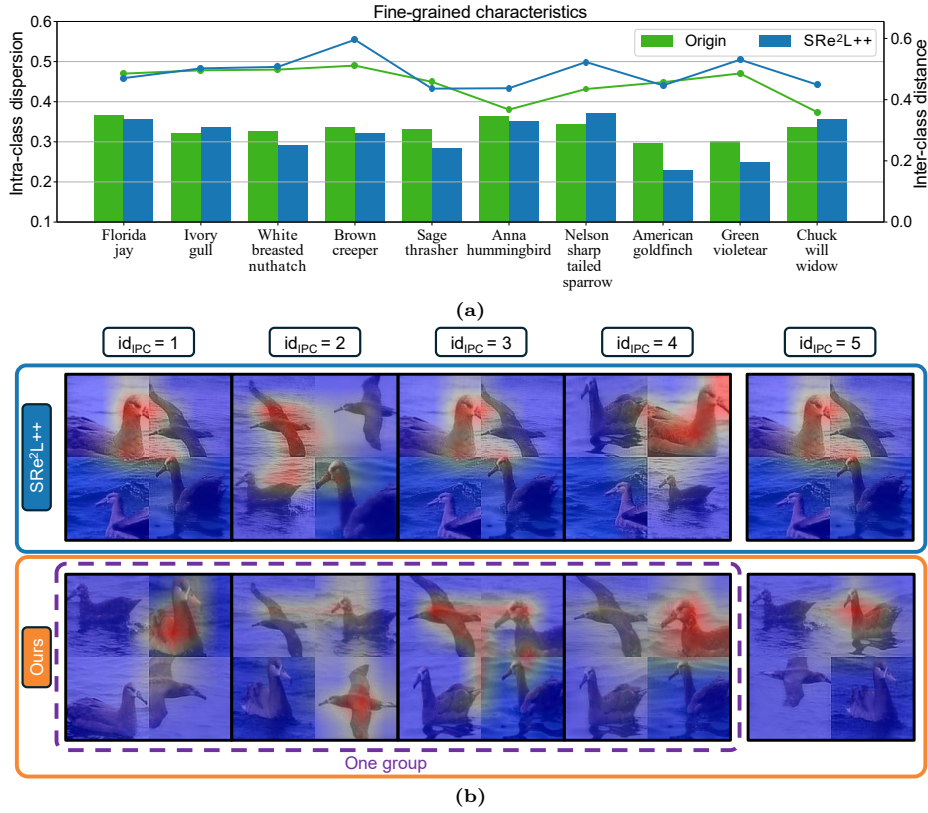


Fig. 1: (a) On CUB-200-2011, 10 classes are randomly sampled to compare the original training set and the SRe²L++ distilled set, where intra-class dispersion (bars) and inter-class distance (line) are reported. (b) Attention heatmaps of *Black-footed Albatross* in CUB-200-2011, comparing the attended regions of distilled samples produced by SRe²L++ and FD².

a compact synthetic dataset that preserves the training utility of the original data, such that a model trained on the distilled set achieves performance close to training on the full dataset while using far fewer resources [5, 16, 18, 39, 46]. Existing DD methods can be broadly grouped into matching-based approaches, *e.g.*, gradient matching [45, 46], distribution matching [20, 36, 48], trajectory matching [1, 6, 17], decoupled methods [3, 4, 43], and generative methods [10, 19, 33, 47]. Most of these techniques are evaluated on general-purpose benchmarks, including ImageNette and ImageWoof [7], where they have shown promising performance. Among these, decoupled methods [3, 4, 43] split dataset distillation into three stages: model pretraining, sample distillation, and soft-label generation, which significantly improves efficiency while preserving competitive accuracy on general benchmarks. However, when directly applied to fine-grained datasets (*e.g.*, CUB-200-2011), existing decoupled methods suffer from two limitations: **(i)** the distilled set preserves unfavorable fine-grained characteristics, and **(ii)** distilled samples within the same class become overly similar.

For (i), Fig. 1a visualizes 10 randomly selected CUB-200-2011 classes, where the original data exhibit large intra-class variation and subtle inter-class differences, which we term the fine-grained characteristic. This characteristic degrades the model’s ability to extract discriminative cues, thereby limiting fine-grained recognition performance. A representative decoupled method, SRe²L++ [4], largely inherits this unfavorable characteristic in the distilled set, and in some classes, the characteristic is even worse than that of the full dataset, making it difficult for the student model to learn discriminative representations. We attribute this to the fact that decoupled DD relies primarily on coarse class-label supervision during distillation: it enforces class-level semantic consistency but does not explicitly encourage intra-class compactness or inter-class separability at the fine-grained level. For (ii), the attention heatmaps in Fig. 1b show that SRe²L++ attends to highly consistent regions across same-class distilled samples, which reduces coverage of discriminative parts and weakens the diversity of discriminative cues available for the student model. This behavior is induced by sample-wise iterative synthesis under nearly identical optimization across iterations, which drives the same-class samples to converge to similar solutions. Together, these issues limit the quality of fine-grained distilled data and consequently degrade the student model’s performance on fine-grained recognition.

To address these issues, we propose **FD²**, a dedicated framework for **F**ine-grained **D**ataset **D**istillation that augments decoupled DD with fine-grained supervision. FD² leverages Counterfactual Attention Learning (CAL) to extract attention maps that localize discriminative regions and to construct class prototypes summarizing these representations. These attention-based signals provide fine-grained supervision for the subsequent distillation process. During pretraining, CAL is used to obtain attention maps and update class prototypes that capture discriminative cues. During distillation, the attention maps are used to localize discriminative regions and are fused with backbone features to form fine-grained representations. Based on these representations, we introduce two complementary constraints. A *fine-grained characteristic constraint* pulls each distilled sample toward its target-class prototype while pushing it away from other-class prototypes, improving intra-class compactness and inter-class separability. A *similarity constraint* encourages diversity among same-class distilled samples by separating their attention distributions, enabling coverage of different discriminative regions. Both constraints are incorporated into existing distillation objectives as plug-in terms, introducing minimal overhead while preserving the original decoupled distillation pipeline. Experiments on multiple fine-grained and general datasets demonstrate that FD² consistently improves representative decoupled distillation methods.

The main contributions of this paper are summarized as follows:

- We propose FD², a dedicated framework for fine-grained dataset distillation that augments decoupled DD with fine-grained supervision. The proposed method improves the structure of the distilled dataset by simultaneously enhancing inter-class separability and promoting within-class diversity.

- We introduce a fine-grained characteristic constraint to improve intra-class compactness and inter-class separability, and a similarity constraint to diversify attention across same-class distilled samples, thereby providing richer localized discriminative cues.
- Extensive experiments on multiple fine-grained and general datasets demonstrate that FD² can be integrated into representative decoupled methods without modification of their workflow and achieves consistent improvements in most settings, with particularly strong gains on fine-grained datasets.

2 Related Work

Dataset Distillation. Dataset distillation (DD) aims to synthesize a compact set of training samples from a large dataset, substantially reducing storage and training cost while retaining performance close to training on the full data. Existing DD methods can be broadly categorized into gradient matching [15, 22, 45, 46], distribution matching [21, 28, 36, 41, 48], trajectory matching [1, 2, 6, 11], decoupled distillation [3, 30, 31, 40, 42, 43], and generative distillation [10, 32, 47, 49].

Decoupled Dataset Distillation. SRe²L [43] pioneers decoupled distillation by separating model pretraining from sample distillation. SRe²L++ [4] improves robustness via stronger augmentation and batch-specific soft labels. FADRM [3] enriches pixel-space information with multi-scale residual connections (with FADRM+ extending to multi-model distillation). G-VBSM [30] enhances generalization across multiple models, while CDA [42] stabilizes optimization through a curriculum strategy. LPLD [40] revisits the necessity of large-scale soft labels and proposes lighter-weight alternatives. Despite their effectiveness on general-purpose benchmarks, these methods are not designed with fine-grained datasets in mind and do not explicitly address fine-grained separability or within-class diversity in the distilled set.

Fine-grained Image Recognition. Fine-grained image recognition (FGIR) often relies on localized discriminative regions, motivating methods that mine and aggregate part-level cues. Cross-X [24] exploits cross-layer interaction with cross-category constraints. PMG [9] learns multi-granularity cues via progressive training. DP-Net [37] injects learnable positional cues and dynamically aligns them with visual content. CSC-Net [38] improves localization by enforcing class-specific semantic coherency. CAL [27] introduces counterfactual attention learning to enhance attention quality and guide models toward key local cues.

3 Approach

3.1 Preliminaries

Dataset distillation aims to synthesize a compact training set that preserves the training utility of a large labeled dataset. Let $\mathcal{T} = \{(x_i, y_i)\}_{i=1}^N$ denote the

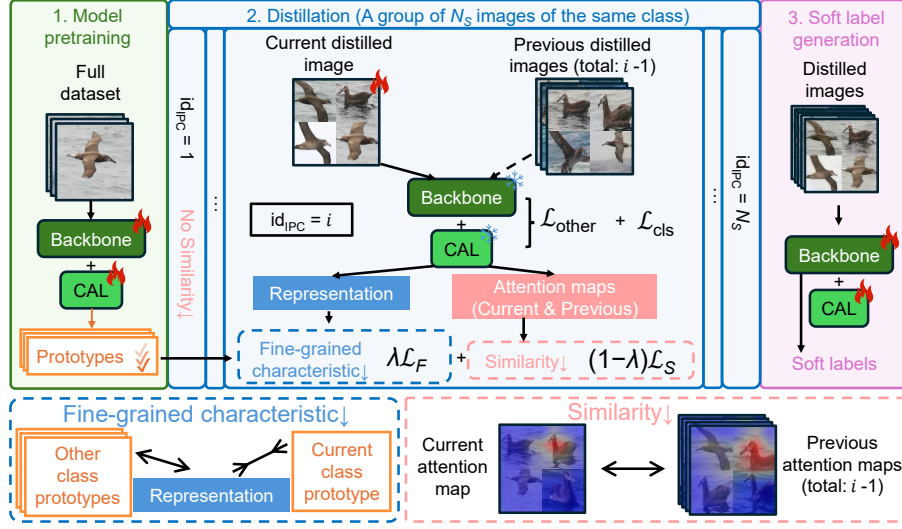


Fig. 2: Overview of FD². 1. We pretrain a Backbone+CAL teacher and maintain class prototypes online. 2. We distill images in the same-class groups of size N_S , adding a fine-grained characteristic constraint (prototype alignment or separation) and a similarity constraint (diverse attention), together with class-supervision from both branches. 3. We generate soft labels using the backbone branch.

original training set. The objective is to distill a substantially smaller set $\mathcal{D} = \{(\tilde{x}_j, \tilde{y}_j)\}_{j=1}^M$ with $M \ll N$, such that models trained on $\mathcal{D}$ achieve performance comparable to those trained on $\mathcal{T}$. This objective can be expressed by minimizing the discrepancy between the losses obtained under the two training regimes:

$$\sup_{(x,y) \sim \mathcal{T}} |\mathcal{L}(f_{\theta_{\mathcal{T}}}(x), y) - \mathcal{L}(f_{\theta_{\mathcal{D}}}(x), y)| \leq \delta, \quad (1)$$

where $\mathcal{L}(\cdot, \cdot)$ denotes the task loss (*e.g.*, cross-entropy), and $\theta_{\mathcal{T}}$ and $\theta_{\mathcal{D}}$ represent the parameters obtained by training on $\mathcal{T}$ and $\mathcal{D}$, respectively. Accordingly, dataset distillation can be formulated as the following optimization problem:

$$\arg \min_{\mathcal{D}, |\mathcal{D}|} \sup_{(x,y) \sim \mathcal{T}} |\mathcal{L}(f_{\theta_{\mathcal{T}}}(x), y) - \mathcal{L}(f_{\theta_{\mathcal{D}}}(x), y)|. \quad (2)$$

Under a given distillation budget, the distilled set $\mathcal{D}$ maintains an equal number of images for each class.

3.2 Overview of FD²

The overall framework of FD² is illustrated in Fig. 2. FD² follows the decoupled distillation paradigm and can be integrated into existing decoupled methods as an additive module. We first pretrain a Backbone+CAL teacher by jointly optimizing the backbone and CAL classifiers while maintaining fine-grained class

prototypes online. During distillation, distilled images are optimized in a sample-wise manner and organized into the same-class groups of size N_S . For each current distilled sample, its feature representation and attention map are obtained from the Backbone+CAL teacher, and the original distillation objective is augmented with two constraints. The fine-grained characteristic constraint pulls the sample toward the prototype of the target class and pushes it away from prototypes of other classes, improving intra-class compactness and inter-class separability. The similarity constraint encourages within-class diversity by increasing the discrepancy between the attention map of the current sample and those of previously generated samples of the same class, thereby promoting coverage of different discriminative regions. Finally, during the soft-label generation stage, soft labels are produced by the backbone branch. Since downstream evaluation trains a standard backbone student, this design helps reduce bias caused by architectural mismatch.

3.3 Counterfactual Attention Learning (CAL)

To address the limitations of previous methods on fine-grained datasets, we introduce Counterfactual Attention Learning (CAL) [27], a simple and effective approach that provides fine-grained class prototypes and precise attention maps. CAL leverages counterfactual intervention to construct a factual-counterfactual contrastive signal, enabling more reliable attention localization and discriminative representation learning under weak supervision.

Given an input image x with label y , a backbone feature extractor $f(\cdot)$ produces the final-stage feature map F , and an attention predictor $g(\cdot)$ generates M attention maps $A = \{A_m\}_{m=1}^M$. CAL then applies an attention-weighted aggregation operator $\Phi(\cdot)$ (*e.g.*, bilinear attention pooling) to obtain the factual representation z :

$$F = f(x), \quad A = g(F), \quad z = \Phi(F, A). \quad (3)$$

To explicitly quantify the contribution of attention, CAL performs a counterfactual intervention by replacing A with an alternative (or perturbed) attention set $\bar{A}$ while keeping F fixed. This operation produces a counterfactual representation $\hat{z}$, and the difference between factual and counterfactual predictions forms an effect prediction:

$$\hat{z} = \Phi(F, \bar{A}), \quad p_{\text{raw}} = Wz, \quad p_{\text{eff}} = p_{\text{raw}} - W\hat{z}, \quad (4)$$

where W denotes the linear classifier of CAL, p_{raw} represents the factual logit, and p_{eff} measures the gain induced by the factual attention relative to the counterfactual attention. Intuitively, if attention focuses on truly discriminative regions, replacing it with $\bar{A}$ weakens class evidence, which increases the discriminability of p_{eff} and is therefore encouraged during training.

Meanwhile, CAL maintains a feature center for each class as a class prototype. Let c_y denote the prototype of class y . The prototype is updated online using a

momentum rule, and a center regularizer encourages the factual representation z to approach the corresponding class prototype:

$$c_y \leftarrow (1 - \mu)c_y + \mu \text{Norm}(z), \quad \mathcal{L}_{\text{center}} = \|z - \text{Norm}(c_y)\|_2^2, \quad (5)$$

where $\text{Norm}(\cdot)$ denotes feature normalization and $\mu \in (0, 1)$ denotes the momentum coefficient.

Subsequently, the pretraining objective of CAL is defined as:

$$\mathcal{L}_{\text{CAL}} = \mathcal{L}_{\text{ce}}(p_{\text{raw}}, y) + \mathcal{L}_{\text{ce}}(p_{\text{eff}}, y) + \eta \mathcal{L}_{\text{center}}, \quad (6)$$

where $\mathcal{L}_{\text{ce}}(\cdot, \cdot)$ denotes the cross-entropy loss and η controls the strength of the center regularization. This objective enables CAL to identify discriminative regions more reliably and to learn discriminative representations during training, which improves classification accuracy.

During the pretraining, distillation, and soft-label generation stages, the teacher adopts the Backbone+CAL architecture. In the post-evaluation stage, following the standard protocol of prior methods, a plain Backbone student model is used. Notably, during pretraining, the backbone classifier and the CAL classifier share the same backbone features and are jointly optimized. This design is motivated by the use of a backbone student model in the post-evaluation stage. If soft labels generated by the CAL branch are used directly, instability may be introduced into the probability outputs of the student model. Formally, the backbone branch predicts:

$$p_{\text{bb}} = W_{\text{bb}} \cdot \text{GAP}(F), \quad (7)$$

where W_{bb} denotes the backbone classifier and $\text{GAP}(\cdot)$ denotes global average pooling. The overall pretraining loss is defined as a weighted combination:

$$\mathcal{L}_{\text{pre}} = (1 - \alpha) \mathcal{L}_{\text{ce}}(p_{\text{bb}}, y) + \alpha \mathcal{L}_{\text{CAL}}, \quad (8)$$

where $\alpha \in [0, 1]$ controls the relative contribution of the two branches. Joint optimization ensures that the backbone teacher model preserves the ability to generate stable soft labels for student training, while still benefiting from the discriminative representations and prototype aggregation provided by CAL.

3.4 Constraints

Suppressing unfavorable fine-grained characteristics and reducing similarity among distilled samples are two key strategies for fine-grained dataset distillation. Therefore, two specialized constraints are introduced.

Fine-grained Characteristic Constraint. Suppressing unfavorable fine-grained characteristics in distilled samples reduces the difficulty for the student model to learn discriminative representations. A straightforward strategy aligns the

representation of a sample with the prototype of its class while increasing the distance to the prototypes of other classes:

$$\mathcal{L}_F(\tilde{x}_{y,i}) = \beta \ell_2(z_{y,i}, c_y) + (1 - \beta) (1 - \mathbb{E}_{k \neq y} [\ell_2(z_{y,i}, c_k)]), \quad \beta \in [0, 1], \quad (9)$$

where $\tilde{x}_{y,i}$ denotes the i -th distilled image of class y in a group, β is a weighting hyperparameter, $\ell_2(u, v) = \|u - v\|_2 / (\|u\|_2 + \|v\|_2 + \varepsilon)$ is a symmetrically normalized Euclidean metric, $z_{y,i}$ denotes the representation of the sample, and c_y and c_k denote the prototypes of the current class and other classes, respectively. According to Eq. (9), the intra-class compactness and inter-class separability of the distilled images are enhanced.

Similarity Constraint. To suppress similarity among distilled samples and encourage diverse discriminative regions, the distance between the attention map of the current sample and the attention maps of previously distilled samples is maximized in the attention space:

$$\mathcal{L}_S(\tilde{x}_{y,i}) = 1 - \mathbb{E}_{j < i} [\ell_2(A_{y,i}, A_{y,j})], \quad 1 < i \leq N_S, \quad (10)$$

and $\mathcal{L}_S$ is omitted when $i = 1$. Here, $A_{y,i}$ denotes the attention map of the current sample, $A_{y,j}$ denotes the attention maps of previous samples of the same class, and N_S denotes the number of distilled images in the group that satisfy this constraint.

Since pretraining adopts a dual-classifier architecture, distillation also supervises class signals with both classifiers to leverage their discriminative capability:

$$\mathcal{L}_{\text{cls}}(\tilde{x}_{y,i}) = (1 - \alpha) \mathcal{L}_{\text{ce}}(p_{\text{bb}}^{y,i}, y) + \alpha \mathcal{L}_{\text{ce}}(p_{\text{cal}}^{y,i}, y), \quad \alpha \in [0, 1]. \quad (11)$$

Here, α is a weighting hyperparameter, and $p_{\text{bb}}^{y,i}$ and $p_{\text{cal}}^{y,i}$ denote the logits produced by the backbone and CAL branches, respectively. Let $\mathcal{L}_{\text{other}}$ denote the original loss of the underlying distillation method. The final objective for optimizing $\tilde{x}_{y,i}$ is

$$\mathcal{L}_{\text{total}} = \mathcal{L}_{\text{other}} + \mathcal{L}_{\text{cls}} + \lambda \mathcal{L}_F + (1 - \lambda) \mathcal{L}_S, \quad (12)$$

where λ controls the relative contribution of the two proposed constraints. The training process is summarized in Algorithm 1. The proposed constraints are jointly optimized with the original objective, which suppresses undesirable fine-grained characteristics and similarity without degrading the effectiveness of the underlying method. A theoretical analysis of these constraints is provided in the supplementary material.

4 Experiments

4.1 Experimental Settings

Datasets. Evaluation of FD² is conducted on three fine-grained datasets, CUB-200-2011 [35], FGVC-Aircraft [26], and Stanford Cars [14], together with two general datasets, ImageNette and ImageWoof, which are subsets of ImageNet-1K [7]. All input images are resized to 224×224 .

Algorithm 1 FD²

Require: Class index y , number of classes K , group size N_S
Require: Pretrained teacher $\phi(\cdot)$ (Backbone+CAL), prototypes $\{c_k\}_{k=1}^K$
Require: Weights α, β, λ , original loss $\mathcal{L}_{\text{other}}$

```

1: for  $i = 1$  to  $N_S$  do
2:   Initialize distilled image  $\tilde{x}_{y,i}$ 
3:   if  $i > 1$  then
4:     for  $j = 1$  to  $i - 1$  do
5:        $(\cdot, \cdot, A_{y,j}) \leftarrow \phi(\tilde{x}_{y,j}) \triangleright$  Attention maps of previous same-class samples
6:     end for
7:   end if
8:   for each step do
9:      $(p_{\text{bb}}^{y,i}, p_{\text{cal}}^{y,i}, z_{y,i}, A_{y,i}) \leftarrow \phi(\tilde{x}_{y,i}) \triangleright$  Logits, representation and attention map
10:    Compute  $\mathcal{L}_F$  by Eq. (9)
11:    if  $i > 1$  then
12:      Compute  $\mathcal{L}_S$  by Eq. (10) using  $\{A_{y,j}\}_{j=1}^{i-1}$  and  $A_{y,i}$ 
13:    else
14:       $\mathcal{L}_S \leftarrow 0$ 
15:    end if
16:    Compute  $\mathcal{L}_{\text{cls}}$  by Eq. (11)
17:     $\mathcal{L}_{\text{total}} \leftarrow \mathcal{L}_{\text{other}} + \mathcal{L}_{\text{cls}} + \lambda \mathcal{L}_F + (1 - \lambda) \mathcal{L}_S$ 
18:    Update  $\tilde{x}_{y,i}$  by gradient descent on  $\mathcal{L}_{\text{total}}$ 
19:  end for
20: end for

```

Comparative Methods. Comparison of FD² is performed with three state-of-the-art DD baselines under the same experimental environment on a single H200 GPU. Since FD² is designed for the decoupled distillation paradigm, SRe²L++ [4] is selected as the primary baseline. This method extends SRe²L [43] through stronger data augmentation and batch-specific soft labels. To evaluate plug-and-play transferability, FADRM+ [3] is also included. FADRM+ is an ensemble-based decoupled method that improves both efficiency and performance through multi-scale residual connections. In addition, RDED [34] is adopted as a representative approach that constructs distilled sets by cropping real images, which can be interpreted as exploiting localized regions during the distillation process.

In cross-architecture generalization experiments and ablation studies, CUB-200-2011 is used as the default dataset unless otherwise specified. The supplementary material provides additional ablation studies, further comparisons with other advanced methods on fine-grained datasets, and efficiency analysis.

4.2 Main Results

Fine-grained Datasets. As shown in Tab. 1, FD² improves the performance of decoupled methods in most settings when compared with SRe²L++ and FADRM+. This result indicates that the mitigation of unfavorable fine-grained characteristics and excessive similarity among samples can improve the quality of distilled data and facilitate the learning of discriminative representations by the

Table 1: Top-1 accuracy is reported in the post-evaluation stage to compare FD^2 with three state-of-the-art methods. For fine-grained datasets, where the number of samples per class is limited, the value of IPC is set to 1, 3, and 5. Following the standard protocol of each method, post-evaluation is conducted for 300 epochs for RDED, which constructs distilled sets by cropping real images, and for 400 epochs for the decoupled baselines SRe^2L++ and $FADRM+$. Evaluation of FD^2 is performed by integrating it into SRe^2L++ and $FADRM+$ (denoted by the subscript “ $_{FD^2}$ ”) while keeping the same post-evaluation settings as the corresponding baseline methods. SRe^2L++ and its FD^2 version follow the standard same-architecture evaluation protocol, while $FADRM+$ and its FD^2 version use multi-teacher for DD and single student for post-evaluation.

Dataset	Student	IPC	RDED	SRe^2L++	$SRe^2L++_{FD^2}$	$FADRM+$	$FADRM+_{FD^2}$
CUB-200-2011	ResNet18	1	38.3	53.4	56.4 ($\uparrow$ 3.0)	54.8	55.0 ($\uparrow$ 0.2)
		3	52.6	60.0	64.9 ($\uparrow$ 4.9)	64.0	64.6 ($\uparrow$ 0.6)
		5	63.9	63.5	67.0 ($\uparrow$ 3.5)	66.4	67.5 ($\uparrow$ 1.1)
	ResNet50	1	33.4	61.1	70.1 ($\uparrow$ 9.0)	66.2	66.5 ($\uparrow$ 0.3)
		3	49.0	65.1	73.7 ($\uparrow$ 8.6)	70.4	70.6 ($\uparrow$ 0.2)
		5	58.6	68.1	75.5 ($\uparrow$ 7.4)	71.5	72.0 ($\uparrow$ 0.5)
FGVC-Aircraft	ResNet18	1	22.1	52.6	58.2 ($\uparrow$ 5.6)	55.0	60.5 ($\uparrow$ 5.5)
		3	36.4	66.6	76.1 ($\uparrow$ 9.5)	72.9	75.1 ($\uparrow$ 2.2)
		5	38.6	68.3	80.0 ($\uparrow$ 11.7)	74.0	77.6 ($\uparrow$ 3.6)
	ResNet50	1	12.8	59.3	67.8 ($\uparrow$ 8.5)	70.8	73.8 ($\uparrow$ 3.0)
		3	34.8	68.0	76.7 ($\uparrow$ 8.7)	76.1	78.9 ($\uparrow$ 2.8)
		5	45.2	72.2	79.0 ($\uparrow$ 6.8)	78.6	81.0 ($\uparrow$ 2.4)
Stanford Cars	ResNet18	1	33.0	52.4	64.5 ($\uparrow$ 12.1)	60.3	74.1 ($\uparrow$ 13.8)
		3	69.4	68.2	75.2 ($\uparrow$ 7.0)	75.0	84.8 ($\uparrow$ 9.8)
		5	76.1	70.9	81.4 ($\uparrow$ 10.5)	77.7	86.6 ($\uparrow$ 8.9)
	ResNet50	1	12.8	65.6	80.7 ($\uparrow$ 15.1)	73.9	85.3 ($\uparrow$ 11.4)
		3	70.2	76.6	86.5 ($\uparrow$ 9.9)	78.9	87.8 ($\uparrow$ 8.9)
		5	75.9	78.0	88.3 ($\uparrow$ 10.3)	82.2	89.0 ($\uparrow$ 6.8)

student model. Notably, the improvements are typically more pronounced when $IPC=1$, with particularly clear gains on the Stanford Cars dataset. For example, with ResNet18, $SRe^2L++_{FD^2}$ exceeds SRe^2L++ by 12.1%, and $FADRM+_{FD^2}$ exceeds $FADRM+$ by 13.8%. This observation suggests that FD^2 can still provide effective discriminative cues even when only a single distilled sample is available for each class. In contrast, RDED exhibits the lowest overall performance, which indicates that coarse cropping of real images does not reliably capture key discriminative regions.

General datasets. To verify the applicability of FD^2 beyond fine-grained datasets, we further evaluate it on ImageNette and ImageWoof. Since SRe^2L++ does not report results on ImageWoof, we use RDED and $FADRM+$ as baselines. As shown in Tab. 2, on ImageWoof (which can be viewed as a 10-class fine-grained subset of ImageNet), $FADRM+_{FD^2}$ achieves clear improvements in most settings, supporting the effectiveness of FD^2 . On ImageNette, the effect of FD^2 depends on the student architecture: with ResNet50, $FADRM+_{FD^2}$ yields clear improvements across all IPC settings, while with the smaller ResNet18, the gains are more limited and become noticeable mainly at $IPC=50$. This suggests that the additional fine-grained cues introduced by FD^2 are more effectively utilized by higher-capacity students. RDED consistently performs the worst across datasets

Table 2: Top-1 accuracy in the post-evaluation stage is reported to compare FD^2 with two baselines. On general datasets, IPC is set to 1, 10, and 50. Following the standard protocols, post-evaluation is conducted for 300 epochs for both RDED and FADRM+. FD^2 is evaluated by integrating it into FADRM+ (denoted by the subscript “ $_{FD^2}$ ”) under the same post-evaluation setting.

Dataset	IPC	ResNet18			ResNet50		
		RDED	FADRM+	FADRM+ $_{FD^2}$	RDED	FADRM+	FADRM+ $_{FD^2}$
ImageNette	1	35.8	39.2	38.6 ($\downarrow$ 0.6)	27.0	31.9	39.0 ($\uparrow$ 7.1)
	10	61.4	69.0	69.5 ($\uparrow$ 0.5)	55.0	68.1	71.4 ($\uparrow$ 3.3)
	50	80.4	84.6	86.7 ($\uparrow$ 2.1)	81.8	85.4	92.3 ($\uparrow$ 6.9)
ImageWoof	1	20.8	22.8	22.1 ($\downarrow$ 0.7)	17.8	19.9	30.3 ($\uparrow$ 10.4)
	10	38.5	57.3	60.7 ($\uparrow$ 3.4)	35.2	54.1	64.6 ($\uparrow$ 10.5)
	50	68.5	72.6	80.0 ($\uparrow$ 7.4)	67.0	71.7	80.2 ($\uparrow$ 8.5)

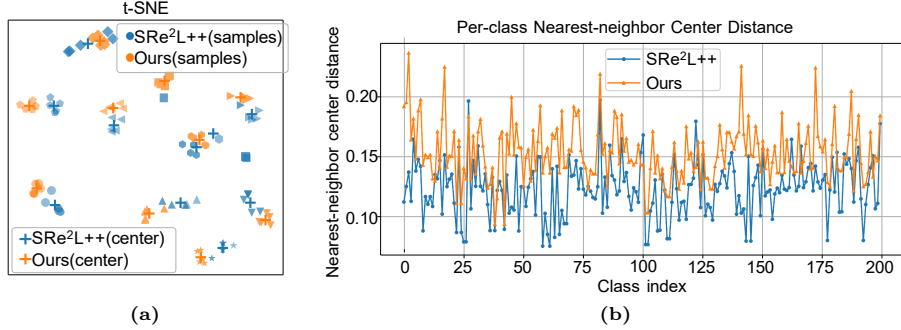


Fig. 3: (a) t-SNE feature distribution. (b) Nearest-neighbor center distance of each class for distilled images on CUB-200-2011.

and IPC settings, indicating that a cropping-based manner provides limited useful information on general datasets. Overall, FD^2 improves performance in most settings, with reduced gains mainly observed for small-capacity students on simpler datasets.

4.3 Analysis

Fine-grained Characteristic. As shown in Fig. 3a, we visualize distilled samples on CUB-200-2011 using t-SNE [25] in the feature space, where 10 classes are randomly sampled with 5 distilled images per class. Compared with SRe²L++, same-class samples distilled with FD^2 exhibit higher intra-class compactness. Meanwhile, Fig. 3b shows larger nearest-neighbor center distances between classes under FD^2 , indicating improved inter-class separability. These results demonstrate that the proposed fine-grained characteristic constraint reduces intra-class variance while enhancing inter-class separability.

Similarity. As illustrated in Fig. 1b, attention heatmaps of the black-footed albatross class in CUB-200-2011 show that SRe²L++ produces highly homogeneous distilled samples within the same class, where the attended regions are nearly

Table 3: Top-1 accuracy for cross-architecture generalization at IPC=3.

Student	SRe ² L++	SRe ² L++ _{FD²}	FADRM+	FADRM+ _{FD²}
ShuffleNetV2 [44]	38.5	47.6 (↑ 9.1)	46.4	46.7 (↑ 0.3)
MobileNetV2 [29]	57.1	62.9 (↑ 5.8)	64.9	65.8 (↑ 0.9)
DenseNet121 [13]	62.0	65.8 (↑ 3.8)	68.6	68.1 (↓ 0.5)
ResNet18 [12]	60.0	64.9 (↑ 4.9)	64.0	64.6 (↑ 0.6)
ResNet50 [12]	64.2	67.7 (↑ 3.5)	70.4	70.6 (↑ 0.2)
ConvNeXt-Tiny [23]	28.9	30.4 (↑ 1.5)	29.1	35.1 (↑ 0.2)

identical across samples. When the similarity constraint is applied within a group of $\text{id}_{\text{IPC}}=1-4$, FD^2 encourages more diverse attention distributions, enriching the discriminative regions in the distilled samples. Although RDED introduces some local information through cropping, this coarse strategy does not reliably capture key discriminative cues. As a result, the performance of RDED is generally lower than that of FD^2 , as shown in Tab. 1 and Tab. 2.

Transferability. FD^2 functions as an add-on module that can be integrated into different decoupled DD methods without modifying the original training pipeline. As shown in Tab. 1, integrating FD^2 into SRe²L++ yields clear accuracy improvements across datasets and IPCs. When integrated into FADRM+, improvements are also observed on most datasets, although the gains on CUB-200-2011 remain relatively limited, while larger improvements appear on Stanford Cars. These results indicate that FD^2 consistently enhances the performance of different decoupled methods and demonstrates strong transferability.

4.4 Cross-Architecture Generalization

Cross-architecture generalization is an important criterion for evaluating the quality of distilled datasets, as it reflects transferability across different students and practical applicability. To examine this property, we compare SRe²L++, FADRM+, and their FD^2 -augmented variants. For SRe²L++ and SRe²L++_{FD²}, we use ResNet18 as the teacher and evaluate with different students; we also include ResNet18 as the student for ease of comparison. For FADRM+ and FADRM+_{FD²}, we distill with multiple teachers and use one student for evaluation. As shown in Tab. 3, integrating FD^2 improves performance for most architectures and methods. On DenseNet121 [13], FADRM+ performs slightly better than FADRM+_{FD²}, while SRe²L++_{FD²} consistently outperforms SRe²L++. These results indicate that FD^2 provides strong cross-architecture generalization and stable performance across diverse evaluation settings.

4.5 Ablation Study

To reduce computational cost and runtime, FD^2 is integrated into SRe²L++, and all ablation studies are conducted using ResNet18 + CAL.

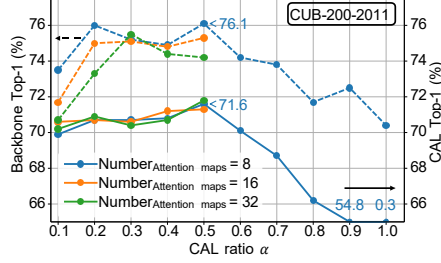


Fig. 4: Top-1 accuracy under different settings during pretraining.

Table 4: Accuracy at IPC= 3 on CUB-200-2011 and FGVC-Aircraft for distilled datasets obtained with different CAL ratios during distillation.

α	CUB-200-2011	FGVC-Aircraft
0.1	62.8	75.1
0.3	63.3	75.9
0.5	63.4	74.8
0.7	62.1	75.6
0.9	62.4	75.2

Impact of the Cal Ratio α in the Model Pretraining and Distillation. Since the backbone classifier and the CAL classifier are jointly optimized, we first perform an ablation study in the pretraining stage to determine an appropriate CAL ratio, as shown in Fig. 4. With the number of attention maps fixed to 8, different CAL ratios are evaluated. When the CAL ratio exceeds 0.5, the accuracy of both the backbone branch and the CAL branch decreases noticeably. Therefore, in experiments with 16 and 32 attention maps, the CAL ratio is restricted to 0.1–0.5. Considering all configurations, the pretrained model with 8 attention maps and a CAL ratio of 0.5 is selected for the distillation stage.

During distillation, the effect of the CAL ratio on post-evaluation performance is further examined with IPC= 3. As shown in Tab. 4, on CUB-200-2011 the highest Top-1 accuracy (63.4%) is obtained at CAL ratio = 0.5. On FGVC-Aircraft, the best accuracy (75.9%) is achieved at CAL ratio = 0.3, which is also optimal in the pretraining stage. These results suggest that using the same CAL ratio in both stages helps maintain consistent feature distributions and attention strengths, thereby reducing stage mismatch. Therefore, adopting the same CAL ratio across stages provides a stable configuration.

Table 5: Impact of β in the fine-grained characteristic constraint.

Method	β	Acc
SRe ² L++ _{FD²}	0.00	63.5
	0.25	63.7
	0.50	64.8
	0.75	64.7
	1.00	64.1

Table 6: Impact of the group size N_S in the similarity constraint.

Method	N_S	Acc
SRe ² L++ _{FD²}	2	64.1
	3	64.9
	4	66.1
	5	64.6

Impact of β in the Fine-grained Characteristic Constraint. To study the relative weighting between alignment with same-class prototypes and repulsion from other-class prototypes, we enable the fine-grained characteristic constraint alone and ablate β at IPC= 3. As shown in Tab. 5, the best post-evaluation accuracy is achieved at $\beta = 0.5$. This result indicates that a balanced weighting improves both intra-class compactness and inter-class separability, thereby improving the quality of samples. Therefore, $\beta = 0.5$ is used in subsequent experiments.

Impact of the Group Size N_S in the Similarity Constraint. To determine an appropriate group size N_S , we ablate N_S with only the similarity constraint enabled and IPC= 5. As shown in Tab. 6, $N_S = 4$ achieves the best post-evaluation accuracy (66.1%). When N_S increases to 5, the available discriminative cues become limited, and additional samples tend to focus on repeated regions. With only a few distilled samples, this repetition may cause the student model to overfit to these regions, which reduces performance (64.6%). Therefore, $N_S = 4$ is selected. For larger IPC, multiple $N_S = 4$ groups can be processed in parallel with multiple processes/GPUs. The distillation times for different N_S settings are reported in the supplementary material.

Table 7: Ablation study of different constraint combinations.

Method	$\mathcal{L}_F$	$\mathcal{L}_S$	Acc
	–	–	63.4
SRe ² L++ _{FD²}	✓	–	64.8
	–	✓	64.6
	✓	✓	64.9

Table 8: The impact of the relative weight λ between the proposed constraints.

Method	λ	Top-1 Acc
	0.2	65.4
SRe ² L++ _{FD²}	0.4	65.0
	0.6	65.7
	0.8	67.0

Effectiveness of the Proposed Constraints. To evaluate the contribution of the proposed constraints, we compare different constraint combinations with IPC= 3. As shown in Tab. 7, both constraints improve performance over the no-constraint baseline (63.4%). The fine-grained characteristic constraint provides a slightly larger gain (64.8%) than the similarity constraint (64.6%). When both constraints are enabled, the highest accuracy (64.9%) is achieved, which confirms their complementary effects.

Effect of the Relative Weight λ between the Proposed Constraints. To determine the trade-off between the proposed constraints, we ablate the relative weight λ under IPC= 5. As shown in Tab. 8, $\lambda = 0.8$ achieves higher post-evaluation Top-1 accuracy, while smaller values (0.2 and 0.4) degrade performance. This suggests that a larger weight for the fine-grained characteristic constraint improves distilled sample quality. Thus, we set $\lambda = 0.8$ in subsequent experiments.

4.6 Visualization

Attention Visualization With and Without Similarity Constraint. As shown in Fig. 5, we compare attention visualizations and mean pairwise cosine similarity (MPCS↓) with and without $\mathcal{L}_S$. Without $\mathcal{L}_S$, same-class samples focus on repeated regions and show higher MPCS; with $\mathcal{L}_S$, they cover richer discriminative regions and show lower MPCS. Thus, although $\mathcal{L}_S$ brings limited accuracy gains in Tab. 7, it improves sample diversity by reducing redundant attention and enriching discriminative cues.

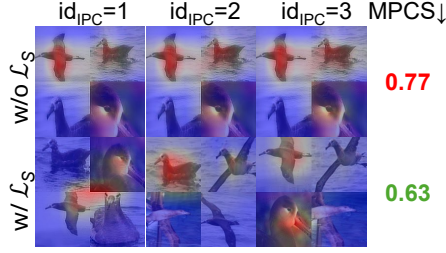


Fig. 5: Attention heatmaps and mean pairwise cosine similarity (MPCS↓) for the $\text{id}_{\text{IPC}} = 1\text{-}3$ distilled samples of the black footed albatross class in CUB-200-2011, with and without the similarity constraint under $N_S = 4$.

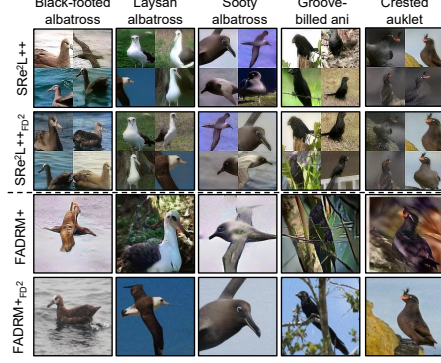


Fig. 6: Visual comparison of distilled samples generated by different methods on five randomly selected classes.

Visualization Comparison of Distilled Samples across Methods. Fig. 6 presents distilled samples from several CUB-200-2011 classes produced by different methods. Samples generated by SReL++ and FADRM+ often exhibit noticeable Gaussian noise. In contrast, integrating FD^2 produces samples with clearer local structures and richer texture details. This observation indicates that the proposed method preserves fine-grained local cues more effectively and improves inter-class separability, which enhances the overall quality of the distilled dataset.

5 Conclusion

We propose FD^2 , a dedicated framework designed for fine-grained dataset distillation. By jointly optimizing a fine-grained characteristic constraint and a similarity constraint together with the original distillation objective, FD^2 mitigates unfavorable fine-grained characteristics and within-class homogenization while preserving the decoupled distillation pipeline. As an add-on module, FD^2 can be seamlessly integrated into decoupled methods and consistently improves performance across diverse fine-grained and general datasets in most settings. Future work will further improve the performance of FD^2 on general datasets and extend it to broader distillation paradigms (*e.g.*, ViT-based models) and larger-scale datasets (*e.g.*, ImageNet-1K).

Acknowledgements

This study was supported by JSPS KAKENHI Grant Numbers JP24K02942 and JP25K21218, and the Data Intelligence Innovation Base Project under Grant No. PPF10120260003 and the Scientific Embodied Multi-Agent System Project under Grant No. PPF10120250008.

References

1. Cazenavette, G., Wang, T., Torralba, A., Efros, A.A., Zhu, J.Y.: Dataset distillation by matching training trajectories. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 10718–10727 (2022) [2](#), [4](#)
2. Chen, M., Huang, B., Lu, J., Li, B., Wang, Y., Cheng, M., Wang, W.: Dataset distillation via adversarial prediction matching. arXiv preprint arXiv:2312.08912 (2023) [4](#)
3. Cui, J., Bi, X., Luo, Y., Zhao, X., Liu, J., Shen, Z.: FADRM: Fast and accurate data residual matching for dataset distillation. In: Proceedings of the Advances in Neural Information Processing Systems (NeurIPS) (2025) [2](#), [4](#), [9](#)
4. Cui, J., Li, Z., Ma, X., Bi, X., Luo, Y., Shen, Z.: Dataset distillation via committee voting. arXiv preprint arXiv:2501.07575 (2025) [2](#), [3](#), [4](#), [9](#)
5. Cui, J., Wang, R., Si, S., Hsieh, C.J.: DC-BENCH: Dataset condensation benchmark. In: Proceedings of the Advances in Neural Information Processing Systems (NeurIPS) (2022) [2](#)
6. Cui, J., Wang, R., Si, S., Hsieh, C.J.: Scaling up dataset distillation to imagenet-1k with constant memory. In: Proceedings of the International Conference on Machine Learning (ICML). pp. 6565–6590 (2023) [2](#), [4](#)
7. Deng, J., Dong, W., Socher, R., Li, L.J., Li, K., Fei-Fei, L.: Imagenet: A large-scale hierarchical image database. In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR). pp. 248–255 (2009) [1](#), [2](#), [8](#)
8. Dosovitskiy, A., Beyer, L., Kolesnikov, A., Weissenborn, D., Zhai, X., Unterthiner, T., Dehghani, M., Minderer, M., Heigold, G., Gelly, S., Uszkoreit, J., Houlsby, N.: An image is worth 16x16 words: Transformers for image recognition at scale. In: Proceedings of the International Conference on Learning Representations (ICLR) (2021) [1](#)
9. Du, R., Chang, D., Bhunia, A.K., Xie, J., Ma, Z., Song, Y.Z., Guo, J.: Fine-grained visual classification via progressive multi-granularity training of jigsaw patches. In: Proceedings of the European Conference on Computer Vision (ECCV). pp. 153–168 (2020) [4](#)
10. Gu, J., Vahidian, S., Kungurtsev, V., Wang, H., Jiang, W., You, Y., Chen, Y.: Efficient dataset distillation via minimax diffusion. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 15793–15803 (2024) [2](#), [4](#)
11. Guo, Z., Wang, K., Cazenavette, G., Li, H., Zhang, K., You, Y.: Towards lossless dataset distillation via difficulty-aligned trajectory matching. In: Proceedings of the International Conference on Learning Representations (ICLR) (2024) [4](#)
12. He, K., Zhang, X., Ren, S., Sun, J.: Deep residual learning for image recognition. In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR). pp. 770–778 (2016) [1](#), [12](#)
13. Huang, G., Liu, Z., Van Der Maaten, L., Weinberger, K.Q.: Densely connected convolutional networks. In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR). pp. 4700–4708 (2017) [12](#)
14. Krause, J., Stark, M., Deng, J., Fei-Fei, L.: 3d object representations for fine-grained categorization. In: Proceedings of the IEEE International Conference on Computer Vision Workshops (ICCVW). pp. 554–561 (2013) [8](#)
15. Lee, S., Chun, S., Jung, S., Yun, S., Yoon, S.: Dataset condensation with contrastive signals. In: Proceedings of the International Conference on Machine Learning (ICML). pp. 12352–12364 (2022) [4](#)

16. Li, G., Togo, R., Ogawa, T., Haseyama, M.: Soft-label anonymous gastric x-ray image distillation. In: Proceedings of the IEEE International Conference on Image Processing (ICIP). pp. 305–309 (2020) [2](#)
17. Li, G., Togo, R., Ogawa, T., Haseyama, M.: Importance-aware adaptive dataset distillation. *Neural Networks* **172**, 106154 (2024) [2](#)
18. Li, G., Zhao, B., Wang, T.: Awesome dataset distillation. <https://github.com/Guang000/Awesome-Dataset-Distillation> (2022) [2](#)
19. Li, L., Li, G., Togo, R., Maeda, K., Ogawa, T., Haseyama, M.: Generative Dataset Distillation: Balancing global structure and local details. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR), Workshop. pp. 7664–7671 (2024) [2](#)
20. Li, W., Li, G., Maeda, K., Ogawa, T., Haseyama, M.: Hyperbolic dataset distillation. In: Proceedings of the Advances in Neural Information Processing Systems (NeurIPS) (2025) [2](#)
21. Liu, S., Wang, K., Yang, X., Ye, J., Wang, X.: Dataset distillation via factorization. In: Proceedings of the Advances in Neural Information Processing Systems (NeurIPS) (2022) [4](#)
22. Liu, Y., Gu, J., Wang, K., Zhu, Z., Jiang, W., You, Y.: DREAM: Efficient dataset distillation by representative matching. In: Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV). pp. 17314–17324 (2023) [4](#)
23. Liu, Z., Mao, H., Wu, C.Y., Feichtenhofer, C., Darrell, T., Xie, S.: A convnet for the 2020s. *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)* (2022) [12](#)
24. Luo, W., Yang, X., Mo, X., Lu, Y., Davis, L.S., Li, J., Yang, J., Lim, S.N.: Cross-x learning for fine-grained visual categorization. In: Proceedings of the IEEE/CVF international conference on computer vision (CVPR). pp. 8242–8251 (2019) [4](#)
25. van der Maaten, L., Hinton, G.: Visualizing data using t-sne. *Journal of Machine Learning Research* pp. 2579–2605 (2008) [11](#)
26. Maji, S., Rahtu, E., Kannala, J., Blaschko, M., Vedaldi, A.: Fine-grained visual classification of aircraft. *arXiv preprint arXiv:1306.5151* (2013) [8](#)
27. Rao, Y., Chen, G., Lu, J., Zhou, J.: Counterfactual attention learning for fine-grained visual categorization and re-identification. In: Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV). pp. 1025–1034 (2021) [4](#), [6](#)
28. Sajedi, A., Khaki, S., Amjadian, E., Liu, L.Z., Lawryshyn, Y.A., Plataniotis, K.N.: DataDAM: Efficient dataset distillation with attention matching. In: Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV). pp. 17097–17107 (2023) [4](#)
29. Sandler, M., Howard, A., Zhu, M., Zhmoginov, A., Chen, L.C.: Mobilenetv2: Inverted residuals and linear bottlenecks. In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR). pp. 4510–4520 (2018) [12](#)
30. Shao, S., Yin, Z., Zhou, M., Zhang, X., Shen, Z.: Generalized large-scale data condensation via various backbone and statistical matching. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 16709–16718 (2024) [4](#)
31. Shao, S., Zhou, Z., Chen, H., Shen, Z.: Elucidating the design space of dataset condensation. In: Proceedings of the Advances in Neural Information Processing Systems (NeurIPS) (2024) [4](#)
32. Su, D., Hou, J., Gao, W., Tian, Y., Tang, B.: D4M: Dataset distillation via disentangled diffusion model. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 5809–5818 (2024) [4](#)

33. Su, D., Hou, J., Li, G., Togo, R., Song, R., Ogawa, T., Haseyama, M.: Generative dataset distillation based on diffusion model. In: Proceedings of the European Conference on Computer Vision (ECCV), Workshop (2024) [2](#)
34. Sun, P., Shi, B., Yu, D., Lin, T.: On the diversity and realism of distilled dataset: An efficient dataset distillation paradigm. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 9390–9399 (2024) [9](#)
35. Wah, C., Branson, S., Welinder, P., Perona, P., Belongie, S.: Caltech-ucsd birds-200-2011. Tech. rep., California Institute of Technology (2011) [8](#)
36. Wang, K., Zhao, B., Peng, X., Zhu, Z., Yang, S., Wang, S., Huang, G., Bilen, H., Wang, X., You, Y.: CAFE: Learning to condense dataset by aligning features. In: Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR). pp. 12196–12205 (2022) [2](#), [4](#)
37. Wang, S., Li, H., Wang, Z., Ouyang, W.: Dynamic position-aware network for fine-grained image recognition. In: Proceedings of the AAAI Conference on Artificial Intelligence (AAAI). pp. 2791–2799 (2021) [4](#)
38. Wang, S., Wang, Z., Li, H., Ouyang, W.: Category-specific semantic coherency learning for fine-grained image recognition. In: Proceedings of the ACM International Conference on Multimedia (ACM MM). pp. 174–183 (2020) [4](#)
39. Wang, T., Zhu, J.Y., Torralba, A., Efros, A.A.: Dataset distillation. arXiv preprint arXiv:1811.10959 (2018) [2](#)
40. Xiao, L., He, Y.: Are large-scale soft labels necessary for large-scale dataset distillation? In: Proceedings of the Advances in Neural Information Processing Systems (NeurIPS) (2024) [4](#)
41. Xue, E., Li, Y., Liu, H., Wang, P., Shen, Y., Wang, H.: Towards adversarially robust dataset distillation by curvature regularization. In: Proceedings of the AAAI Conference on Artificial Intelligence (AAAI) (2025) [4](#)
42. Yin, Z., Shen, Z.: Dataset distillation in large data era. Transactions on Machine Learning Research (2024) [4](#)
43. Yin, Z., Xing, E., Shen, Z.: Squeeze, recover and relabel: Dataset condensation at imagenet scale from a new perspective. In: Proceedings of the Advances in Neural Information Processing Systems (NeurIPS) (2023) [2](#), [4](#), [9](#)
44. Zhang, X., Zhou, X., Lin, M., Sun, J.: Shufflenet: An extremely efficient convolutional neural network for mobile devices. In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR). pp. 6848–6856 (2018) [12](#)
45. Zhao, B., Bilen, H.: Dataset condensation with differentiable siamese augmentation. In: Proceedings of the International Conference on Machine Learning (ICML). pp. 12674–12685 (2021) [2](#), [4](#)
46. Zhao, B., Bilen, H.: Dataset condensation with gradient matching. In: Proceedings of the International Conference on Learning Representations (ICLR) (2021) [2](#), [4](#)
47. Zhao, B., Bilen, H.: Synthesizing informative training samples with gan. In: Proceedings of the Advances in Neural Information Processing Systems (NeurIPS), Workshop (2022) [2](#), [4](#)
48. Zhao, B., Bilen, H.: Dataset condensation with distribution matching. In: Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision (WACV). pp. 6514–6523 (2023) [2](#), [4](#)
49. Zhong, X., Fang, H., Chen, B., Gu, X., Qiu, M., Qi, S., Xia, S.T.: Hierarchical features matter: A deep exploration of progressive parameterization method for dataset distillation. In: Proceedings of the Computer Vision and Pattern Recognition Conference (CVPR). pp. 30462–30471 (2025) [4](#)